

文脈と対象世界モデルを利用した 機械翻訳へ向けて

Steps toward a Machine Translation Using Context and World Model

石 崎 俊^{*1}
Shun Ishizaki

井佐原 均^{*1}
Hitoshi Isahara

徳永 健伸^{*2}
Takenobu Tokunaga

田中 穂積^{*2}
Hozumi Tanaka

- *1 電子技術総合研究所知能情報部自然言語研究室
Natural Language Section, Machine Understanding Division, Electrotechnical Laboratory.
*2 東京工業大学工学部情報工学科
Dept. of Information Eng., Faculty of Eng., Tokyo Institute of Technology.

1989年9月5日 受理

Keywords: context, world model, machine translation, focus model.

1. は じ め に

我が国では既に機械翻訳システムの商品化が行われ、研究開発も盛んである。コンピュータのハードウェアが計算スピードと ULSI 化で象徴されるのに対し、ソフトウェアの高度さは知的な処理（使いやすさ）によって代表されている。そのような知的処理の具体的な例として機械翻訳は絶好の応用分野である。

ところで、機械翻訳システムの中核は自然言語処理技術であり、形態素解析、構文解析、意味解析、文脈解析、文章生成技術を組み合わせたものである。従来の自然言語処理研究では、文法規則を中心にした構文解析や浅い意味解析が主な研究対象となっている。したがって、構文情報などだけでは曖昧で、適切な訳語が決定できない場合が多く、システムへの入力の前処理・前編集や、システムの出力を後で修正する後編集にかなり人手が必要となっている。このように技術的に不十分な点を埋め、人による前後の処理を減らすには、コンピュータが入力文の構文や意味内容を的確に把握し、正しく出力文に反映させなければならない。しかし、実は、これは大変難しい技術を含んでいる⁽¹⁾⁻⁽³⁾。

文章の意味内容の解析・理解には、浅い理解から深い理解へといくつかの段階が考えられている。単語レベルの浅い理解は形容詞と名詞の係り受けなどがあり、文レベルの理解は、例えば、“take a bus”と

“take a photo”を「バスに乗る」と「写真を撮る」に訳し分けるように、一つの文中で処理される解析である。

文脈レベルでは、パラグラフ程度の複数の文からなる文章を対象とし、パラグラフ全体で表現しているトピックを大局的に把握し、そのトピックに関して我々人間の持っている常識を用いて、従来の浅い理解では曖昧で適切な意味がわからなかった場合も知的に処理しようとする解析法である。

対象世界モデルとは、解析する文章の意味内容が、あるトピック（話題）について述べているときに、そのトピックについて人間が（常識として）標準的に持っている知識をモデル化したものである⁽³⁾⁻⁽⁶⁾。

また、談話構造の特徴として、焦点の処理、照応解析、さらに会話文などにおける話者の意図の理解なども文脈レベルと考えられる⁽⁶⁾⁻⁽¹⁰⁾。

文脈レベルよりもさらに深いものとして、言語行為や比喩のように、言外の意味を推論する必要のあるレベルも考えられるが、文脈レベルとの境界が明確に引かれているわけではない。

ここで、日本語と英語を考えると、両者は大きく異なる言語であるため、翻訳のためには浅い解析では不十分であることが多い。英語がしっかりとした文法を持つのに比べて、日本語では省略が多く英語ほどリジッドではない。したがって、解析して得る構文木の変換を主として行うアプローチでは、十分な品質の翻

訳を行うことは難しい。文文法だけでなく談話文法の研究の重要性が指摘されている⁽¹¹⁾。

翻訳システムが有用であるためには、文レベルで二つの言語を対応づけるだけでなく段落（意味的なまとまりのあるいくつかの文）レベルの情報を利用する必要がある。例えば、日本語と英語では段落構造が大きく異なるから、有用なシステムを作るには文脈情報が必要になる。

ところで、「文脈」の解釈や定義には言語学、心理学などの分野でもさまざまなものがあり、我々の工学の中でも専門家によって使い方が異なる場合が多い。

本稿では、“context”の訳として従来から親しみのある「文脈」を使うことにする。文脈には一般に次のような定義がある⁽¹²⁾。

- (1) 広義：状況（situation）と同じ意味に用いられ、その環境や場面などの条件を表す言語外的なものも含む。
- (2) 狭義：① 文章の中での文と文との続きぐあい、照応、省略解析。
② 文中での意味の続きぐあい。

本稿で採用する文脈の定義は、上記の(1)広義を中心として、(2)狭義の①照応、省略解析を含むものである。また、照応、省略解析を行う上で、(1)広義の文脈が必要不可欠であると考えられる。

以下の2章では、文脈情報を使う機械翻訳システムの一つである CONTRAST の簡単な紹介を行う。3章では、Sidner の焦点モデルを中心に、訳語選択における文脈情報の役割を論ずる。4章では、2章と3章で述べた機械翻訳技術の課題と将来動向について述べる。

2. 文脈情報を用いる機械翻訳

我が国における機械翻訳の研究は、1960年頃から長期間にわたって続けられており、いくつかの国家的プロジェクトもあって多くの成果を得ている⁽¹²⁾。欧米でも同様な歴史を持ち、ECでは膨大なニーズに対処すべく研究が進められている⁽¹³⁾。

ここで、文脈情報を用いる機械翻訳システムの一例として、現在研究開発中の簡単なプロトタイプ of CONTRAST (CONtext TRAnslation) を紹介する。

CONTRAST の解析と生成部分では図1にあるように、文法情報、語彙情報、文脈情報を使用する。文法や語彙情報と同様に、概念辞書に格納した文脈レベルの情報を含む背景知識を用いて入力文章を理解する⁽¹⁴⁾⁽¹⁵⁾。

入力文章は、通常一つのトピックについてのテキストであり、中間表現（理解結果）は、そのトピックについての首尾一貫した木構造となる。生成システムは中間表現に基づいて、出力文の段落構造、段落内の文章構造、各文の文体などを決定する⁽¹⁶⁾⁽¹⁷⁾。

CONTRAST は、Schank が提唱するようなプリミティブ概念⁽¹⁸⁾にすべて解析するわけではない。それは解析手続きが複雑になり、推論の量が多すぎて、実際の機械翻訳システムを構築することは難しいのである。したがって、この CONTRAST で採用する概念表現における概念の粒度（granularity、概念体系で上位概念ほど粒度が粗く、下位ほど細かい）は、あまり粗いものではないので、語彙項目から概念への変換辞書には、多くの場合に対応する概念が用意されている。

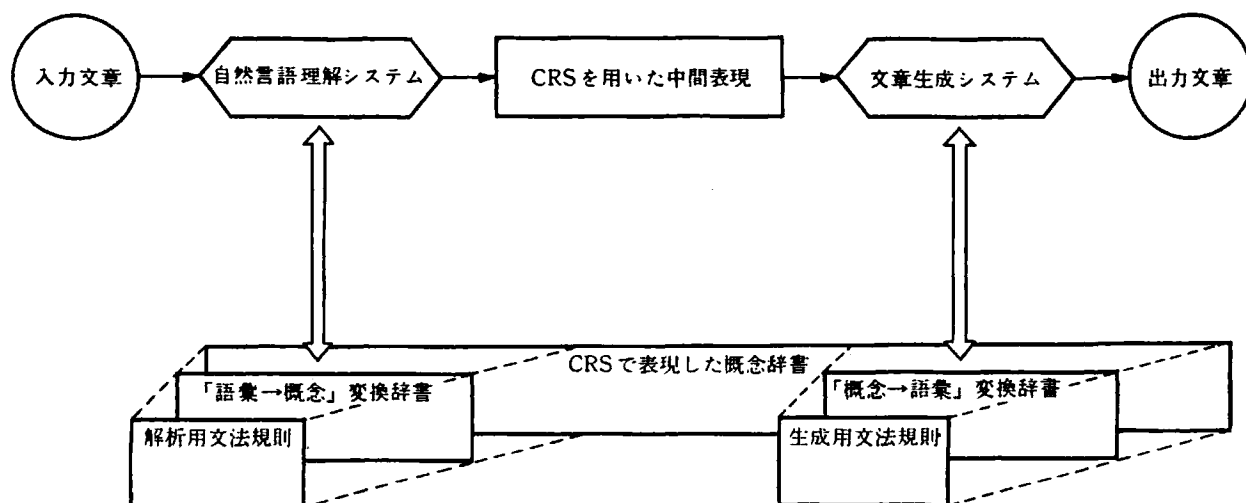


図1 文脈情報機械翻訳システム CONTRAST の構成
(CRS: 文脈表現構造, Contextual Representation Structure)

2.1 文脈情報のための知識表現⁽¹⁹⁾

(1) 概念とは

CONTRAST は, CRS (Contextual Representation Structure) という文脈情報を考慮した概念表現構造を用いている。CRS は処理対象言語に応じて導入されるさまざまな概念を互いに関連づけることにより, 多言語の処理を統一的に取り扱おうとするものである。この構造を構成する要素としては, 大きく分けて「概念」と各概念に結びついた「インスタンス」の二つがある。概念はインスタンスの集合としての特性の集まりであると定義する。本節では, 説明上の混乱を避けるために, 概念は「*」で始まる表記(「*弟」)とし, インスタンスは「@」で始まる表記(「@弟」)として, 言語の表層の単語(「弟」)と区別する。

概念は予め作られた概念辞書上で上位概念・下位概念による階層構造をなしており, 各インスタンスの持つべきスロットや部分構造, 制約が規定されている⁽²⁰⁾。

(2) 概念の表現形式

概念辞書の書式を図2に示す。概念辞書の記述の中で, GENERALIZATIONS (上位概念)とSPECIFICATIONS (下位概念)とは概念の階層構造を示すために用いられる。SPECIFICATIONS には下位概念が並べられていて, 互いに背反な概念どうしは括弧でくくられている。SLOTS-OF-INSTANCE には, この概念から導出されるインスタンスが持つ可能性のあるスロット (AGENT, OBJECT, LOCATION, TIME などの関係概念)とスロットを埋める値についての制約条件が記述される。スロットの情報は概念の階層構造に沿って継承される。したがって上位概念で定義されたスロットは, より詳細な定義をする必要がない限り, 下位概念では定義する必要がない。

SLOTS-OF-INSTANCE の中で, キーワード“SCENES”に記述された情報は, この概念(のインスタンス)が時間関係あるいは因果関係によってどのような概念(のインスタンス)に分割されるかが示されている。また, キーワード“PART-OF”を用いれば, この概念の部分構造を示すことができる。なお, SCENES は時間軸による概念の部分構造化であるから, 本来, PART-OF に含まれるものである。

(3) インスタンスの表現

以上述べてきたような概念体系から, 実際の文章を解析する過程で順次インスタンスが導出されていき, 全体の文脈表現構造が抽出される。この過程で, 事象や物に関する一つの対象をいくつかの概念のインスタ

```
(概念 (GENERALIZATIONS 概念1 概念2 ...)
(SPECIFICATIONS 概念3 概念4 ...)
(SLOTS-OF-INSTANCE
  (スロット1 条件1)
  (スロット2 条件2)
  ...
(SCENES シーン1 シーン2 ...)
(PART-OF パート1 パート2 ...))
(CONSTRAINTS ...))
```

図2 概念辞書の書式

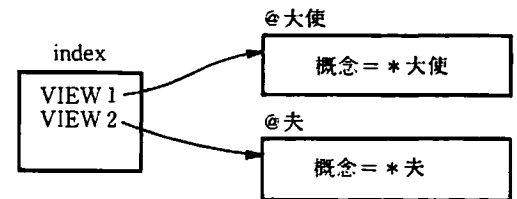


図3 「大使でありかつ夫であるもの」の表現形式

ンスの組合せで表現する必要が生じる場合がある。ここで「対象」とは文脈表現構造が表現しようとするものであり, 物概念においては, その外延となる。ここでは, 複数のインスタンスのそれぞれがその対象を別の視点 (view) から説明しているものと考え, 一つの対象に対する視点をまとめて指す節点として index という節点を導入した。ここで, 視点 (view) は複数のインスタンスの組合せで表現されるような対象の持つ情報を, 各インスタンスごとに切り出したものである。また, index は, 視点の一部あるいは全部を指す節点であり, これにより対象の一部を特定して言及することが可能になる。外部からの対象の参照は, index あるいは view に対して行われる。「大使であり, かつ夫であるもの」を示す表現を図3に示す。ここで, 「@夫」と「@大使」とは, ある外延をそれぞれの視点から説明していることになり, それらすべてを指す index の節点が, すべての視点からの情報を含んだ外延全体を示している。

(4) 関係概念について

概念には大きく分類して物概念と事象概念とがある。事象概念は動作概念と状態概念に分かれる。状態概念はさらに関係概念と呼ばれるものを含み, このインスタンスは概念のインスタンスどうしの関係を定義する(図4)。「*動作主 (agent)」や「*対象 (object)」といったスロットも概念の一つとして定義されている。関係概念の中で, このように三項関係として定義されるものは, そのインスタンスが結びつける二つの概念のうちのいずれかにスロットとして埋め込むことが可能である。

概念を規定するには, 上下関係の構造を示すだけでは十分でない。上下関係がないその他の概念も含めた

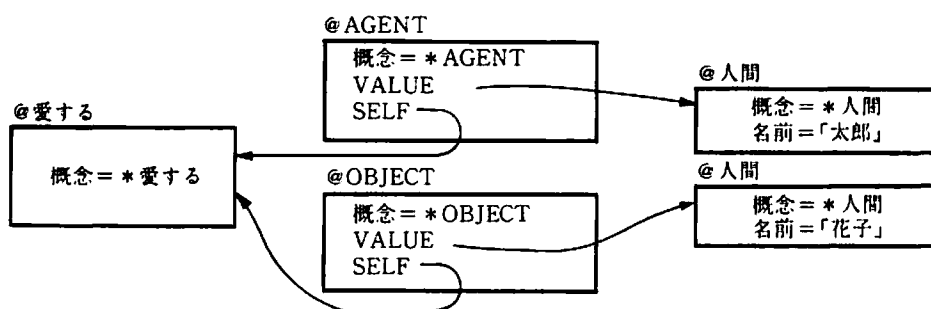


図4 関係概念 (AGENT, OBJECT) のインスタンスによる表現法

各種の関係によってその概念を定義するために「制約」を用いる。例えば、運転手という概念を記述するための制約として、乗物を運転するという動作の動作主であると記述する。

2.2 文脈理解

(1) 構文解析・意味解析部の概要

CONTRAST で用いているパーザは、単語辞書を文法規則と同様の表現法で内部表現することにより、単語の切り出し過程においても、横型の探索を可能とした完全横型探索を行うものである⁽²¹⁾。右方枝分かれ構造の構文解析木を導出するような文法規則からは特別な場合を除いて、ただ一つの解析結果しか導出されない。

このような、係り受けに関する曖昧性を内に含んだ解析結果を導出することは、もとの文における曖昧性をこの時点では解決せずに、決定を「保留」しておくことになる⁽²²⁾。

CONTRAST の意味解析部は、拡張 LINGOL 上で開発されていた意味解析システム EXPLUS⁽²³⁾を発展させたものである。右方枝分かれ構造の構文解析木上で、意味解析を進めることにより、意味情報による構文的曖昧性の除去および係り受けの非交差の法則の自然な実現が可能となった。

図5は入力した新聞記事の一例である。例えば、「ペドロ・マヌエル・デアリステギ」は「氏」によって人名であるとみなされ、全体として「人間」という概念のインスタンスとなる。これと「スペイン大使」という概念との間で、係り受けの可能性を調べることにより意味解析を行う。すなわち、「スペイン大使」中のスロットを「人間」が占めうるか、あるいは、「人間」中のスロットを「大使」が占めうるかが調べられる。この場合、「人間」の「役職」として、「大使」が埋め込まれることになる。続いて、ここで得られた「人間」のインスタンスと「誘拐する」とが、助詞「が」の情報を合わせて関係づけられる。この場合、「人間」が

誘拐3時間
無事に解放
スペイン大使
【ベイルート10日＝共同】西
ベイルートで10日午後、スハイ
ン大使のペドロ・マヌエル・デ
アリステギ氏が誘拐されたが、
事件発生後三時間たった同夜無
事解放された。
スペイン大使館によれば、大
使は犯人の二人組に目隠しをさ
れたまま、車で二、三ほど運ば
れ、ビル内に監禁された。しか
し、酒を出されたうえ本を与え
られるなど扱いは一変だったこ
う。

図5 入力した新聞記事の一例

「誘拐する」の「対象」であるとみなされる。このようにして図6に示す意味解析結果を得る。

図5の第1パラグラフで、「西ベイルートで10日午後」という時空位置を表す表現は、その後のすべての動詞句に係ることができる。しかし、「開放する」には「同夜」という時間情報があるため、「10日午後」がこれに係ることはできない。一方、「西ベイルート」という場所情報は二つの動詞句に係りうる。

「同夜」の「同」については、システムは「同夜」が出現する前方向へ探索を行い、日付を示す情報「10日」を得る。この結果、TIME73 (同夜) のDATE スロットに10が入り、TIME71 と TIME73 が同じ日の午後と夜であることが示される。

(2) 文脈解析部の概要

(a) 入力文の話題の決定⁽²⁴⁾

人間が図5に示すような新聞記事から、そのトピックに合致した記憶構造 (先験的文脈表現構造) を決定するためには、記事のタイトル (サブタイトルも含む) が重要な役割を果たす。このような考え方から作られたのが、「タイトル中のキーワードによる方法」である。

タイトル中に特定の単語 (名詞および動詞) が存在した場合には、それに対応する文脈表現構造が選択される。実際このような手法ではほとんどの場合に適切な文脈表現構造が選択される。これはある意味では当然のことである。新聞記事は、タイトルを読んだだけでかなりの情報を得られるように通常はタイトルが決められている。

入力情報の同化は、既存文脈表現構造を用いて演繹的に入力情報を理解する（文脈表現構造の適切な位置に格納する）ことに対応する。

入力情報の同化には次の4種類がある。

- ① 並列的事象：一つのシーンにいくつかの事象が並列に同化される。例えば、爆発による被害のシーンに負傷を表す事象や、物の破壊を表す事象が並列的に同化される。
- ② 既出事象の詳細化：同一の事象を数回に分けて述べる場合で、2回目以降は既にシーンに埋め込まれている事象の詳細化であり、同一事象として同化される。
- ③ 文脈表現構造中の未だ埋まっていないシーンに対応する事象：この事象は新たに対応するシーンに埋め込まれる。
- ④ 文脈表現構造中のシーンに対応しないが、それらと因果関係によって結合する事象：因果関係などを表す関係概念を用いて既出事象に結びつけられる。

同化のための判定基準は上記①～④に対応して、

- ① 一つのシーンの下位概念に含まれ、かつ、既に入れた入力事象に含まれなければ、そのシーン中の並列事象
- ② 既出事象と同一概念で、それぞれの持つ付加情報に明らかな矛盾がなければ、詳細化
- ③ これら①、②に該当せず文脈表現構造の他のシーンに含まれる事象
- ④ 文脈表現構造中のシーンに含まれず、入力文中の因果関係で既出事象と結びついた事象

というものである。

(c) アクティブマッチングを用いた文脈理解⁽²⁴⁾

意味解析、文脈解析の過程では、指示物の同一性の検出が大きな課題となる。言語表現の表層上でいくつかの表現があって、それらの指示物を同一のものとして同定するとき、異なるスロットにさまざまな情報が格納されていても、同一スロットに異なる値があるような明確な矛盾がないかぎり、積極的に同一物としてマッチングを取る手法を採用しており、これを「アクティブマッチング法」と名付けた。これは我々人間が、文章を読んだり、人の話を聞いて理解するとき、大変効率的に理解が行われることの一つの基準を示すものである。

中間表現は、トピックを表す適切な文脈表現構造に入力新聞記事から得られる知識をボトムアップに埋め込んで得られたものである。図7のAとOで表されたものは、それぞれ図の上部にある構造を指す。このA

とO、すなわち、加害者である「犯人」と被害者である「大使」とが、表層では明示的に表現されていないような事象（例えば、“release”）において、それぞれ“agent”と“object”になるということの決定は、上に述べたアクティブマッチングによる。すなわち、予めシーン構造中に規定してあるシーン動作のスロットどうしの対応関係に基づいて決定する。この例のように、“release”と上位の“kidnap”とのスロット値が同じである場合は、defaultとして明示しない。また、“convey”の“situation”は、“blindfold（目隠しをする）”によって引き起こされる状態のもとで“convey”が行われたことを、入力文から抽出し表現している。

2.3 文章生成⁽¹⁶⁾

本システムの英文生成の特徴は、構造化された意味表現からの生成であることである。生成すべきテキストに対応する意味表現は、書き手の思考内容に相当すると考えられる。また、構造化された背景知識である概念辞書に、書き手の常識に相当するレベルでの判断の基礎を与えている。本生成システムでは、この文脈表現構造の構造的性質を生かし、その構造自体の持つ情報を生成に使用している。

生成システムは中間表現から出発し、概念辞書も参照しながら英文法に従って目標言語によるテキスト（本研究の場合は英語による新聞記事）を生成するが、そこで、英語の文法や単語についての知識ばかりでなく、英語による実際のテキストでよく用いられる全体構成や、主語選択、テンスやアスペクトの選択などに関する知識も使用する必要がある。

また、これまであまり検討されてこなかった複数パラグラフのテキストの生成を試みている。パラグラフ構成は言語や分野による違いもあり、広く受け入れられる理論は見当たらない。内容的にまとまりのあるテキストを生成するには、テキスト全体の首尾一貫性の考察が不可欠である。パラグラフどうしの内容的な連関性が、意味表現の構造的連関性に基づくことを指摘しテキストの意味表現構造と首尾一貫性との関係を理論化し、文脈表現構造からの首尾一貫性のある複数パラグラフからなるテキストの生成規則に結びつけている。

生成部は、文脈表現構造、概念辞書、その概念辞書を英語に変換する部分、英文生成規則からなる英語辞書を用いる。詳しい生成規則などは文献(16)を参照されたい。

このようにして生成した英語文章を図8に示す。同図の英文の名詞句の生成における人名および指示代名詞の使用法では、パラグラフ構造に依存した手法を考

A SPANISH AMBASSADOR WAS KIDNAPPED BY 2 CRIMINALS IN THE AFTERNOON ON THE 10TH IN WESTBEIRUT AND HE WAS RELEASED SAFE BY THEM ON THE NIGHT 3 HOURS LATER .

THE AMBASSADOR PEDORO-MANUERU-DEARISUTEGI WAS TAKEN BLINDFOLDED BY THE CRIMINALS BY CAR ABOUT 2 KILOMETERS AND HE WAS CONFINED BY THEM INSIDE OF A BUILDING .

図 8 英文生成出力の例(図5の一部に対応)

案している。また、このシステムでは、a, an, the など定性のルール化を進めているが、日時の表現法なども含めてまだ検討の必要がある。

3. 訳語選択における文脈情報の役割

3.1 はじめに

本章では、問題を代名詞の照応の解決に絞り、照応の解決が訳語選択に及ぼす影響について考察し、実際に英日翻訳の実験システムを構築することによってその効果を検証する。照応解析には、Sidner の焦点モデル⁽⁹⁾⁽¹⁰⁾を用い、できるだけ浅い解析で照応を決定する。

3.2 照応と訳語選択

この節では、代名詞の照応の解決が訳語選択に及ぼす影響について具体例を挙げて考察する。

英語を日本語に翻訳する場合、代名詞は翻訳されずに省略されるか、その代名詞が指示している名詞が明示的に繰返し用いられることが多い⁽²⁸⁾。文を翻訳の単位として扱う翻訳システムでは、代名詞の翻訳は目標言語の対応する適当な代名詞に翻訳し、実際にそれが何を指すのかは、出力された文章を読む人間の推論能力にまかせるしかない。例えば、

- ① Taro bought a new camera.

(太郎は新しいカメラを買った。)

- ② He likes it.

(彼はそれが気に入っている。)

において、「彼」が太郎であり、「それ」がカメラであることは、文を単位として翻訳している限り認識できない。例文②では、「he」、「it」をそれぞれ「彼」、「それ」と翻訳しても人間が読めば、「彼」、「それ」が何を指しているかがわかるので、それほど不自然な訳とはならない。しかしながら、代名詞が何を指示しているかによって、格関係やその他の語の訳語が大きく影響を受ける場合がある。例えば、

- ③ Taro bought a new camera.

(太郎は新しいカメラを買った。)

- ④ He had Hanako take his picture with it.

(彼はそれ(カメラ)で花子に写真を撮ってもらった。)

- ⑤ Taro had a little dog.

(太郎は小さい犬を飼っていた。)

- ⑥ He had Hanako take his picture with it.

(彼はそれ(犬)と一緒に花子に写真を撮ってもらった。)

例文④と例文⑥は全く同じだが、先行する文によって「it」の指示する対象が異なる。例文④では、「it」がカメラを指示しているため、「it」は主動詞「take」の道具格となるが、例文⑥では、「it」が犬を指示しているので随伴格になる。この例では、指示対象によって格関係が異なるため訳が全く異なってしまう。このような場合は、代名詞「it」を単に「それ」と訳したのでは不十分で、照応関係を正しく解析しないと翻訳できない。

もうひとつの例を挙げよう。

- ⑦ Taro bought a film.

(太郎はフィルムを買った。)

- ⑧ He took a picture.

(彼は写真を撮った。)

- ⑨ He developed it.

(彼はそれ(フィルム)を現像した。)

この例では、例文⑨の「it」が何を指示しているかによって、主動詞「develop」の訳を選択しなければならない。もし、先行する文脈から「it」が指示するものが、例えば、「camera」などであつたら、この「develop」は「開発する」と訳すほうが適当である。

Carter によれば、このような代名詞の照応の多くは浅い解析で解決できるので⁽⁷⁾、文脈情報の利用の第一段階として照応処理を翻訳システムに組み込むことは有意義なことである。

3.3 照応処理を含む翻訳システム

本節では、Sidner の焦点モデルを用いて、照応処理を行う機構を組み込んだ翻訳システムについて述べる。

(1) Sidner の焦点モデル

Sidner は焦点を用いて局所的な情報から照応を解決するモデルを提案している⁽⁹⁾⁽¹⁰⁾。このモデルでは、

話者が談話中で中心に置く要素(焦点)は、照応詞の先行詞になりやすいという仮定がなされている。以下、Sidner の焦点モデルを簡単に説明する。

Sidner のモデルでは、以下の六つのレジスタを用いて焦点の状態を表す。

- ・ DF (Discourse Focus), AF (Actor Focus)
現在、焦点となっている要素が入る。
- ・ PDF (Potential Discourse Focus), PAF (Potential Actor Focus)
最後に読まれた文で焦点になれなかった要素が入る。
- ・ DFS (Discourse Focus Stack), AFS (Actor Focus Stack)
焦点の移動が起こったとき、古い焦点を積む。

照応の解析アルゴリズムは以下ようになる。

- ③ 最初の文に対して焦点候補を DF に設定する。
- ④ 焦点を使って照応の解析を行う。
- ⑤ 解決した照応により、焦点の状態(各レジスタの値)を更新する。

③は最初の文にのみ適用され、④と⑤は文が読み込まれるたびに適用される。

焦点候補は次の優先順位によって決定する。

- ③ 主語 (be 動詞構文, there 挿入文)
- ④ それ以外の構文
 - ① 対象格
 - ② 動作主格を除く要素
 - ③ 動作主格
 - ④ 動詞句

なお、焦点になれなかった要素は PDF にリストとして設定される。AF, PAF, AFS はレジスタの内容が動作主になりうる要素であるという点を除き、DF, PDF, DFS と同様に操作される。

(2) システムの概要

本節では Sidner の焦点モデルを組み込んだ、実験的な機械翻訳システムについて述べる。本システムは LangLAB⁽²⁵⁾ をベースに実装されており、Prolog で記述されている。入力された文章は LangLAB 上で 1 文ずつ構文解析、意味解析が行われるが、同時に照応詞が含まれる場合は先行詞の決定も行う。照応詞の解決は照応詞の意味処理を行う時点で起動する。1 文が処理されると、(1)で述べた Sidner のアルゴリズムを適用して焦点の管理を行う。解析の結果として文の意味を表す中間表現が得られる。したがって、この段階で多義語の訳語選択は解決されていることになる(ただし、このシステムでは多訳語については考慮していない)。この中間表現からトップダウンに出力文

を生成する。

例として、例文⑦～⑨について解析の様子を見てみよう。以下に、各文が解析されたあとの中間表現と焦点レジスタの内容を示す。中間表現の中で*がついているものは、照応の解決によって決まったものであることを示している。

“Taro bought a film.”

DF :film(expected) AF :taro(expected)
PDF :buy PAF :[]
PFS :[] AFS :[]

中間表現:

[物を買う, AGENT:[太郎]]

OBJECT:[フィルム SPEC:[不定]]

“He took a picture.”

DF :film(confirmed) AF :taro(confirmed)
PDF :picture, take PAF :[]
DFS :[] AFS :[]

中間表現:

[写真などを撮る, AGENT:[*太郎]]

OBJECT:[写真, SPEC:[不定]]

“He developed it.”

DF :film AF :taro
PDF :develop PAF :[]
DFS :[] AFS :[]

中間表現:

[現像する, AGENT:[*太郎]]

OBJECT:[フィルム]]

3.4 ま と め

本章では正しい翻訳を得るためには 1 文内の情報だけではなく、文脈情報も利用する必要があることを指摘し、照応の解決と翻訳との関係および Sidner の焦点モデルによる照応処理を組み込んだ機械翻訳の実験システムについて述べた。

Sidner の焦点モデルでは、局所的な文脈情報に基づいて照応を決定するために処理が比較的軽いという利点があるが、反面、より大域的な情報を考慮する必要があったり、常識的なより深い推論を必要とする場合は照応が解決できない。例えば、

⑩ He took a picture.

(彼は写真を撮った。)

⑪ He developed it.

(彼はそれ(フィルム)を現像した。)

の場合、例文⑦、⑨と異なり、“it”の先行詞が文章中に明示的に現れていないので、「写真を撮る」という動作に「フィルム」が関係し、「写真を撮ったら普通フィルムを現像する」という常識的な知識を用いて深い推論を行わないと照応を解決することができない。対象とする領域を予め限定して大域的な談話構造をモデル化する研究はあるが（例えば文献（8）など）、領域を限定しないモデルについては今後の課題である。

本稿では述べなかったが、翻訳システムを考える場合、高品質の訳を得るためには、照応の解決だけでなく生成における照応詞の使用も重要な問題である。原言語で代名詞となっていたものを目標言語でも代名詞で表現しなければならないというわけではない。特に、日本語の場合は代名詞化されるより省略される場合が多い。人間によって翻訳された文章などを調べ、原言語と目標言語の照応詞の使い方などを調査する必要がある。

4. 次世代機械翻訳技術

従来の用語法における中間言語方式とは、対象とするいくつかの言語に共通の概念の体系を作り、入力文の持つ情報をすべてそのような中間言語で表現してから生成するという方法として考えられてきた。しかも本論文で述べてきたような文脈レベルの解析や生成技術が不十分な状況で、そのような方式を考えるのだから、難しい課題が山積していることがわかる。

中間表現を表すための言語として、いわゆる 에스เปรント語や英語、日本語などの言葉を考えていない。それらの言葉は状況意味論でいうところの「効率」の良い言葉である。すなわち、一つの単語が文脈によっていくつもの意味を持つことができるのである。これを文脈依存性と呼ぶが、本論文の文脈解析技術はこのような文脈依存性を解消するための研究であり、文脈生成技術は逆に、文脈に依存した効率の良い文章を生成する研究である。

したがって、このような文脈処理技術を用いれば、中間表現として、普通の言葉（英語、日本語など）を用いる必要性が減少する。2・1節で述べた概念は、そのような意味で曖昧性が少なく、一通りの意味に対応させている^{(27) (28)}。

このような概念の総体は対象とする各言語を覆うものである。ここでいう概念の世界は、各対象言語から抽出される概念の集合の AND をとった共通部分ではなく、むしろ OR をとった総合的なものである⁽²⁷⁾。この意味で、各国語に固有の概念も内包している。

ところが、もし概念の世界の中で、各国語の語彙に対応する概念を言語ごとに分離していたら、従来のトランスファ方式と変わらないものになってしまう。

文脈情報を用いる方式では、このような概念の世界がいかに全体として体系づけられているかが問題であり、そのように体系づけられた概念と言語の間を文脈処理技術によって継ぐことが中心的な課題の一つといえよう。

対象世界モデルとしての概念の世界を構築し、概念全体の体系化と個々の概念記述における文脈レベルの情報のあり方について研究を進める必要がある。

文脈の解析と生成は、いろいろな知識を使いながら深い処理をするため、領域固有性がある。一つの文を処理するにも多くの知識を必要としている。

また、概念の記述形式における制約記述は、宣言的に書かれており、ユニフィケーションによって処理されるという観点から、制約プログラミングに馴染むものと考えられる。そのようなプログラミング手法によって効率的に処理する可能性がある^{(29) (30)}。

一方、多言語翻訳を考えると、目標言語に該当する言葉がない場合の処理が課題である。CONTRASTでは、そのために概念展開法を研究している⁽³¹⁾。これは、生成システムが中間表現中の概念に対して目標言語の語彙が見いだせないときに起動される。そのとき、概念はまず概念階層における細かい粒度の詳しい記述を持つ下位概念に展開される。もしそこで、目標言語の語彙に対応づけられた概念が発見できないときは、逆に、粒度の粗い上位概念をたどって近似的な訳を作り出すことが考えられる。

5. 今後の研究

常識を用いて深い推論を行う大域的な談話構造のモデルによる照応処理の研究が今後必要であるが、このような大域的モデルの領域限定性の問題がある。また、照応詞の使用における言語ごとの綿密な調査もこれからの重要な研究であろう。

意図の理解は会話の解析・生成・翻訳で大切なテーマである。発話者の意図の解析は、プラン・ゴール⁽⁴⁾などの大域的な知識と、発話文から得られるボトムアップな知識の統合・同化を行う必要がある。

比較的深い理解を伴う翻訳方式は、トランスファ方式に比べて情報が減るという意見がある。しかし、むしろ逆であると思われる。文脈などの深い理解を行うためには概念辞書を使う。概念辞書に、談話構造なども含めてどのくらいの質と量の情報が書かれている

か、それをどのくらいうまく使えるかによって情報が
増えるか減るかが決まる。自然言語の解析と生成に十
分使えるような概念辞書、すなわち、特定領域の知識
や常識を組み込んだ概念辞書というものを構築し、そ
れを使って翻訳することによって情報が増え、正確な
翻訳による高品質性が得られるようになるとされる
(32)(33)。ただし、いまの技術レベルで直ちにシス
テム化し商品化するという意味ではなく、これからの
方向性としてそのような課題で研究を進めるという意
味である。また、コンピュータのメモリーやスピード
の性能などの発展や、自然言語処理に適した論理型コ
ンピュータの今後の発展に期待したい。利用可能な知

識の質と量、推論の深さが、それらにも依存してい
るからである。

謝 辞

電総研知能情報部の中島部長をはじめ、自然言語研
究室の皆さん、東工大田中研究室の皆さんには有効な
助言、議論などでお世話になったので感謝したい。
CONTRAST 研究グループの内田ユリ子氏(英語生
成)、橋田浩一氏(文脈情報のための知識表現、現在
ICOT)は著者たちと共同で本論文の研究をされた。
ここに謝意を表したい。

◇ 参 考 文 献 ◇

- (1) 田中穂積：機械翻訳の誤りと人間の誤り，津田塾会設立
40周年記念日本語国際シンポジウム「日本語教育の現代的
課題」，pp. 177-186，(財)津田塾会(1988)。
- (2) 田中穂積：言語理解研究の諸相，人工知能学会誌，Vol. 3，
No. 3，pp. 271-279(1988)。
- (3) 石崎 俊，井佐原 均：文脈処理技術，情報処理，27 巻 8
号，pp. 897-905(1986)。
- (4) Schank, R. C. and Abelson, R. : Script, Plans, Goals
and Understanding, LEA (1977)。
- (5) Schank, R. C. : Dynamic Memory, Cambridge Uni-
versity Press (1982)。
- (6) Allen, J. : Natural Language Understanding, The
Benjamin/Cummings Publishing Co., Inc. (1987)。
- (7) Carter, D. : Interpreting Anaphors in Natural Lan-
guage Texts, Ellis Norwood LTD. (1987)。
- (8) Grosz, B. : The representation and use of focus in a
system for understanding dialogues, IJCAI '77, pp.
67-76(1977)。
- (9) Sidner, C. L. : Focusing in the comprehension of
definite anaphora, in Brady, M. and Berwick, R. C.
eds., Computational Models of Discourse, pp. 267-
330, MIT Press (1983)。
- (10) Sidner, C. L. : Toward a Computational Theory of
Definite Anaphora Comprehension in English Dis-
course, MIT, AI Lab., TR-537 (1979)。
- (11) 井上和子：文法から談話文法へ，言語，12 巻 12 号，pp.
38-46(1983)。
- (12) 長尾 真，ほか：科学技術庁機械翻訳プロジェクトの概要，
情報処理，26 巻 10 号，pp. 1203-1213(1985)。
- (13) Slocum, J. (ed.) : Special Issues on Machine Trans-
lation, Computational Linguistics, Vol. 11, Nos. 1-3
(1985)。
- (14) Ishizaki, S., Isahara, H. and Handa, K. : Natural
Language Processing System with Deductive Learn-
ing Mechanism, in Nagao, M. ed, Language and AI,
pp. 125-144, Elsevier Science Publishers B. V. (1987)。
- (15) Isahara, H. and Ishizaki, S. : Context Analysis Sys-
tem for Japanese Text, 11th International Conference
on Computational Linguistics (1986)。
- (16) 内田ユリ子，石崎 俊，井佐原 均：実証的な知識に基づ
く文脈表現構造からの英語テキスト生成，電子情報通信学
会論文誌，Vol. 72-D-II, No. 9(1989)。
- (17) Ishizaki, S. : Generating Japanese Text from Con-
ceptual Representation, in McDonald, D. and Bolc,
L. eds., Natural Language Generation Systems, pp.
256-279, Springer-Verlag (1988)。
- (18) Schank, R. C. : Conceptual Information Processing,
North Holland (1975)。
- (19) 石崎 俊，井佐原 均，橋田浩一，内田ユリ子，横山晶一：
文脈理解のための概念記述法，情報処理学会自然言語処理
研究会 NL64-7(1987)。
- (20) 井佐原 均，橋田浩一，内田ユリ子，石崎 俊：文脈解析
における概念照合のための推論と制約記述について，情報
処理学会第 35 回全国大会(1987)。
- (21) 元吉文男，井佐原 均，石崎 俊：日本語用完全構文探索
構文解析法，情報処理学会第 32 回全国大会(1986)。
- (22) 井佐原 均，石崎 俊：日本語新聞記事解析における構文
情報および意味情報の抽出法，人工知能学会誌，Vol. 3，
No. 5，pp. 607-616(1988)。
- (23) 田中穂積：計算機による自然言語の意味処理に関する研
究，電子技術総合研究所研究報告第 797 号(1979)。
- (24) 石崎 俊，井佐原 均：文脈と言語理解，電子通信学会技
術研究報告，NLC86-4(1986)。
- (25) 徳永健伸，田中穂積：視点を考慮した概念の同一化，情報
処理学会第 37 回全国大会，pp. 1284-1285(1988)。
- (26) 樺垣 実：日英比較表現論，大修館書店(1975)。
- (27) 石崎 俊，内田裕士：多言語間翻訳のための中間言語につ
いて，情報処理学会自然言語処理研究会，NL70-3(1989)。
- (28) 石崎 俊：コンピュータで文間を補う一意味と文脈一，数
理科学，No. 309(1989. 3)。
- (29) 石崎 俊，浮田輝彦，田村直良，橋田浩一：意味・言語・
対話，コンピュータソフトウェア，Vol. 6, No. 4(1989)。
- (30) 向井国昭：談話理解とロジック，人工知能学会誌，Vol. 3，
No. 3，pp. 289-300(1988)。
- (31) 井佐原 均，橋田浩一，内田ユリ子，石崎 俊：機械翻訳
システム CONTRAST における概念変換について，情報
処理学会第 34 回全国大会(1987)。
- (32) Ishizaki, S., Sakamoto, Y., Ikeda, T. and Isahara,
H. : Machine Translation Systems Developed at Elec-
trotechnical Laboratory, Future Computing Sys-
tems, Vol. 2, No. 3, pp. 275-297(1989)。
- (33) 石崎 俊，井佐原 均：文脈情報翻訳システム CON-
TRAST，情報処理，Vol. 30, No. 10(1989)。

著者紹介



石崎 俊 (正会員)

1970年東京大学工学部計数工学科卒業。同学科助手を経て、1972年電子技術総合研究所入所。1981～82年イエール大学客員研究員。1984年パターン情報部推論システム研究室長。現在、知能情報部自然言語研究室長、工学博士。自然言語処理、学習、音声処理の研究に従事。ACL、AAAI、電子情報通信学会、日本認知科学会、情報処理学会等各会員。



徳永 健伸 (正会員)

1983年東京工業大学工学部情報工学科卒業。1985年同大学院理工学研究科修士課程修了。同年(株)三菱総合研究所入社。1986年東京工業大学大学院博士課程入学。1987年より同大学工学部情報工学科助手。自然言語処理、知識表現に関する研究に従事。情報処理学会、日本ソフトウェア科学会、認知科学会各会員。



井佐原 均

1978年京都大学工学部電気工学第二学科卒業。1980年同大学院工学研究科電気工学専攻修士課程修了。同年、電子技術総合研究所入所。現在、知能情報部自然言語研究室主任研究官。自然言語理解、知識表現の研究に従事。情報処理学会、日本認知科学会、ACL各会員。



田中 穂積 (正会員)

1964年東京工業大学理工学部制御工学科卒業。1966年同大学院修士課程修了。同年電気試験所(現、電子技術総合研究所)入所。1983年東京工業大学工学部情報工学科助教授。1986年同大学教授。工学博士。人工知能、自然言語処理の研究に従事。情報処理学会、電子情報通信学会、認知科学会、計量国語学会各会員。